

Article

Comparison of Racog and Racog-Rus for Classifying Imbalanced Data on Gradient Boosting and Naïve Bayes Performance

Rahmi Fadhilah¹, Heri Kuswanto^{1,*}, Dedy Dwi Prastyo¹, Dinda Ayu Safira¹
and Muhammad Yahya Matdoan^{1,2}

¹ Department of Statistics, Institut Teknologi Sepuluh Nopember, Surabaya 60111, Indonesia

² Department of Statistics, Pattimura University, Ambon 97233, Indonesia

* Correspondence: heri.kuswanto@its.ac.id

How To Cite: Fadhilah, R.; Kuswanto, H.; Prastyo, D.D.; et al. Comparison of Racog and Racog-Rus for Classifying Imbalanced Data on Gradient Boosting and Naïve Bayes Performance. *Journal of Modern Applied Statistical Methods* **2025**, *24*(1), 6. <https://doi.org/10.56801/Jmasm.V24.i1.6>

Abstract: This study aims to determine the effect of resampling RACOG and RACOG-RUS data on Gradient Boosting and Naïve Bayes classification in predicting water quality with unbalanced data. The data used in this study were 720 data from January 2022 to December 2023. It was found that Gradient Boosting performed best when using RACOG-RUS resampling data and feature selection with a number of numInstances of 200. While Naïve Bayes has the best performance when using RACOG-RUS resampling data without feature selection with a number of numInstances of 300. It can be seen that resampling RACOG data does not outperform RACOG-RUS in both classification models because it is known that the data generated in RACOG does not make the dataset more balanced than RACOG-RUS. Hybrid sampling is necessary if RACOG samples are used as the training dataset.

Keywords: gradient boosting; naive bayes; RACOG; RACOG-RUS

1. Introduction

River water is a source of life that plays an important role in human life and the surrounding ecosystem. In addition, rivers have an important function as a transportation system that carries large amounts of materials in dissolved and particulate form from natural and artificial sources in one direction. However, river water often becomes a reservoir for harmful industrial waste, causing pollution that endangers human health and the surrounding environment. Human activities related to agricultural pesticide use and changes in land use are major factors affecting surface water quality. The impacts are devastating, with millions of deaths of living beings every year and huge economic losses. Therefore, effective protection and management of river water are critical to maintaining global sustainability in both life and the economy [1,2].

Class imbalance in water quality data is a common problem, where one outcome is less common than the other. Although standard classifiers such as logistic regression, SVM, and decision tree are effective with balanced data, they are suboptimal in the case of imbalance [3,4]. Imbalanced data occurs when there is an imbalance of classes during classification, resulting in majority and minority classes. This results in misclassification in the majority class more often than the minority class. Thus, the focus should be on reducing misclassification in the majority class and sacrificing accuracy in the minority class [5]. The solution to the problem of data imbalance can be divided into three main approaches. The first approach, data-level methods, is an external approach that focuses on balancing the data by reducing samples from the majority class (undersampling) or adding samples to the minority class (oversampling). The second algorithm-level method is an internal approach aimed at addressing bias due to data imbalance through improving existing algorithms or developing new classification algorithms.



The third method, the hybrid method, is a combination of data-level and algorithm-level methods with the aim of improving classification accuracy in the face of data imbalance problems [6].

Various approaches have been proposed to handle the imbalanced data problem, including data-level techniques such as oversampling and undersampling. Oversampling techniques used such as Rapidly Converging Gibbs Sampler (RACOG) [7] and undersampling techniques used such as Random Under Sampling (RUS) [8]. Recent developments combine these two techniques in hybrid sampling to create balanced datasets. However, inadequate representation of minority classes and overlap issues between classes can compromise model performance. Therefore, an effective hybrid method is needed to address these issues by considering the probability distribution of minority classes and avoiding over-sampling. The hybrid sampling hybrid sampling technique performed by [9] used RACOG-RUS.

The ensemble approach and cost-sensitive learning are two common methods used in handling data imbalance in machine learning studies [6]. The ensemble method combines the decisions of multiple base classifiers to produce more accurate predictions, where the diversity and accuracy of each base classifier become key factors in good performance. Three popular ensemble methods are boosting, bagging, and stacking [6,10]. In addition, cost-sensitive learning allocates costs to different classes to improve model performance. However, classification results can be unstable with this approach due to the difficulty in determining the appropriate error cost [6]. Although several models have been developed to predict water pollution, the development of unbalanced learning systems with established methods is still insufficient in water quality studies. In order to improve the performance of the model in the case of high-dimensional models, feature selection is used.

Gradient boosting is an extension of the concept of ensemble learning, where several weak prediction models are combined to form a stronger model. This method extends the idea of boosting by taking into account the gradient of the loss function when adding new prediction models to the ensemble. As a machine learning method, gradient boosting is used for supervised learning applications, such as classification and regression. It uses an ensemble of weak prediction models, usually a decision tree, and has three main components: loss function, weak learner, and additive [11]. Gradient boosting has advantages in handling high-order relationships in data and various data challenges [12]. In addition, Naive Bayes is a machine learning algorithm that uses Bayes' Theorem for classification. Known for its high training speed, it works under the assumption of independence between features. It is used for high-dimensional datasets and generates probability predictions for each class by utilizing the joint posterior probability distribution between classes and attributes. Despite relying on fairly simple assumptions, Naive Bayes can compete well with other algorithms in many cases [13,14]. A performance comparison between Gradient Boosting and Naive Bayes in water quality research can provide important insights into the best approach for specific situations.

Based on previous research that has been described, this study combines data imbalance handling with RACOG and RACOG-RUS applied to the gradient boosting and naive bayes methods with and without feature selection on water quality data in the Bengawan Solo River using various evaluations, including accuracy, balanced accuracy, sensitivity, specificity, precision, F-measure, and AUC.

2. Feature Selection

The dimensionality of data used in machine learning tasks can lead to significant issues, such as the curse of dimensionality in existing learning methods. Feature selection is one widely used solution to reduce dimensionality. Its objective is to remove specific subsets of irrelevant features from the original dataset based on evaluation criteria of importance and to select a few features that are most representative of the original set. Feature selection typically results in improved learning performance, reduced computational costs, and enhanced model interpretation. Methods for feature selection can be classified into filter, embedded, and wrapper methods. Filter methods use general data characteristics to assess features without requiring a classifier in the process. Meanwhile, wrapper methods rely on the accuracy of specific classifiers to select features, and embedded methods integrate feature subsets as internal mechanisms in the classifier training process. This study employs two methods for feature selection: wrapper and embedded methods. In the wrapper method, the Boruta feature selection algorithm is chosen. The Boruta algorithm is a wrapper method based on random forest algorithms that capture important and interesting features in the dataset while considering the output variable. To support this decision, an embedded method is also used for feature selection. This method utilizes the important features of the classifier as an internal mechanism to measure the predictive strength of each feature and selects the highest predictive strength. This method is chosen because it can provide a simple approach to feature selection by using the average accuracy and an average decrease in node impurity.

3. Rapidly Converging Gibbs Sampler (RACOG)

RACOG is a data resampling method using Gibbs sampling to generate new minority class samples from the minority class probability distribution approximated using the Chow-Liu algorithm with the Markov Chain Monte Carlo (MCMC) method [7]. RACOG offers an alternative mechanism for selecting the initial value of the random variable X (denoted as $X^{(0)}$) to improve the randomly selected Standard Gibbs sampling. In addition, it uses the minority class data points as the initial sample set and runs the Gibbs Sampler for each minority class sample. The total number of iterations for the Gibbs Sampler is limited according to the distribution of the minority and majority classes. RACOG generates multiple Markov chains, each starting with a different minority class sample, unlike the conventional Gibbs sampler, which starts with a minority class sample in one very long chain. [15] said that this approach is different from other approaches because it considers the distribution of minority classes and is proven to provide the best performance when compared to other oversampling approaches. MCMC is an increasingly popular method for obtaining information about distributions, especially for estimating posterior distributions in Bayesian inference [16]. The following is the RACOG algorithm with time complexity where n = data dimension, D = minority class cardinality, and T = predetermined number of iterations, as presented in Algorithm 1.

Algorithm 1. RACOG

```

1: function RACOG (minority,  $D$ ,  $n$ ,  $\beta$ ,  $\alpha$ ,  $T$ )
Input: minority = minority class data points;  $D$  = size of minority ;  $n$  =
        minority dimensions ;  $\beta$  = burn-in period;  $\alpha$  = lag;  $T$  = total number of
        iterations
Output: new_samples = new minority class samples
2. Construct Dependence tree  $DT$  using Chow-Liu algorithm.
3. for  $d = 1$  to
    $D$  do
4.    $X^{(0)} = \text{minority}(d)$ 
   for  $t = 1$ 
     to  $T$ 
5.   do
6.     for  $i = 1$  to  $n$  do
       Simplify
7.        $P(X_i | x_1^{(t+1)}, \dots, x_{i-1}^{(t+1)}, x_{i+1}^{(t)}, \dots, x_n^{(t)})$ 
       using  $DT$ 
        $x_i^{(t+1)} \sim P(S_i)$  where  $S_i$  is the state space of
8.       attribute  $x_i$ 
       if  $t > \beta$  AND  $t \bmod \alpha = 0$ 
          $\text{new\_samples} = \text{new\_samples} + X(t)$ 
return
   $\text{new\_samples}$ 

```

4. RACOG-RUS

The RACOG method is known to simply improve the traditional Gibbs Sampler to address class imbalance without considering the usefulness of the resulting samples. As a result, there is a risk that RACOG may add unnecessary minority class samples, potentially causing overfitting issues and providing little help in building good hypotheses [7,17]. Therefore, [18] proposed a new approach that combines the RACOG model with undersampling techniques (RUS) in RACOG-RUS to eliminate redundant minority class samples. In this method, RACOG uses Gibbs Sampler to synthesize the minority classes estimated using the Markov Chain Monte Carlo (MCMC) method. In this algorithm, each sampling step considers the univariate conditional distribution of each dimension. The value of each dimension depends on the values of other dimensions and the previous values of the same dimension. Such conditional distributions are easier to model compared to the complete joint distribution. The standard Gibbs Sampler algorithm is presented in Algorithm 2.

Algorithm 2. Gibbs Sampler

```

1.  $X^{(0)} = \langle x_1^{(0)}, \dots, x_n^{(0)} \rangle$ 
   for  $t$ 
2.    $= 1$  to  $T$ 
3.   for  $i = 1$  to  $n$ 
4.      $x_i^{(t+1)} \sim P(X_i | x_1^{(t+1)}, \dots, x_{i-1}^{(t+1)}, x_{i+1}^{(t)}, \dots, x_n^{(t)})$ 

```

$P(X_i | x_1^{(t+1)}, \dots, x_{i-1}^{(t+1)}, x_{i+1}^{(t)}, \dots, x_n^{(t)})$ is drawn from the joint probability distribution that represents the univariate conditional distribution of each dimension for selecting attribute values, thus generating a new sample of the minority class $x_i^{(t+1)}$ which is determined by randomly selecting from the state space distribution with all possible values of the attribute x_i . Gibbs Sampling is implemented with two factors, namely burn-in and lag. Burn-in refers to the number of iterations required to generate samples in order to achieve a stable distribution. Lag, on the other hand, pertains to successive samples that are removed from the Markov Chain after each sample is received to prevent correlation between subsequent samples. After oversampling from the minority class, there is a reduction of samples from the majority class to align the data to its original size, thus providing more accurate information. RACOG-RUS combines probabilistic oversampling while considering the minority class distribution and undersampling to produce a more balanced and representative dataset while overcoming the problem of overfitting.

5. Gradient Boosting

Gradient boosting (GBM) is one of the machine learning methods used in supervised machine learning applications, and it includes various classification and regression problems. GBM builds a prediction model in the form of a collection of weak prediction models, such as decision trees. GBM consists of three main components: loss function for optimization, weak learner for prediction, and additive model to combine the weak learner to optimize the loss function [11,19]. This algorithm differs from Random Forests (RF) in that it trains weak learners sequentially, where each new model attempts to correct the errors made by the ensemble of previous models. The general procedure of GBM is described in Algorithm 3. As a non-linear method, gradient boosting naturally outperforms linear models when higher-order relationships exist in the data and has demonstrated its superiority compared to other machine learning algorithms in various data challenges [12].

$$g_t(x) = E_y \left[\frac{\partial \psi(y, f(x))}{\partial f(x)} | x \right]_{f(x)=\hat{f}_{t-1}(x)} \quad (1)$$

$$(\rho_t, \theta_t) = \arg \min_{\rho, \theta} \sum_{i=1}^N [-g_t(x_i) + \rho h(x_i, \theta)]^2 \quad (2)$$

Algorithm 3. Gradient Boost Algorithm**Inputs:**

- input data $(x, y)_{i=1}^N$
- number of iterations M
- choice of the loss-function $\psi(y, f)$
- choice of the base-learner model $h(x, \theta)$

Algorithm:

```

1: initialize  $\hat{f}_0$  with a constant
2: for  $t = 1$  to  $M$  do
3:   compute the negative gradient  $g_t(x)$ 
4:   fit a new base-learner function  $h(x, \theta_t)$ 
5:   find the best gradient descent step-size  $\rho_t$ :
      $\rho_t = \arg \min_{\rho} \sum_{i=1}^N \psi[y_i, \hat{f}_{t-1}(x_i) + \rho h(x_i, \theta_t)]$  for  $i = 1$  to  $n$  do
6:   update the function estimate:
      $\hat{f}_t \leftarrow \hat{f}_{t-1} + \rho_t h(x, \theta_t)$ 
7: end for

```

6. Naïve Bayes

A naive Bayes classifier is a machine learning algorithm that utilizes Bayes' theorem to solve classification problems by taking probabilities into account. It belongs to the supervised learning category and is known for its high training speed. Naive Bayes operates on the assumption of independence between predictors by evaluating each feature separately and collecting simple statistics per class of each feature. The advantage of this algorithm is its suitability for high-dimensional datasets, as it relies on the assumption of statistical independence among the features. By utilizing the joint posterior probability distribution between classes and attributes, Naive Bayes can effectively derive classification problems by generating probability predictions for each class. Research shows that Naive Bayes is able to compete well with other machine learning algorithms, even outperforming them in some cases [13]. The algorithm relies on the assumption that conditional attribute values are independent of each other, taking into account the target value of the given instance [2,8].

$$P(a_1, a_2, a_3, \dots, a_n | v) = \arg \max \prod_{i=1}^n P(a_i | v_j) \quad (3)$$

In other words, the probability of observing a conjunction of attribute values given a target instance is simply the product of the probabilities for the individual attributes. Although this assumption is sometimes unrealistic in real contexts, Naive Bayes remains effective due to its ability to simplify probability calculations and reduce dimensionality. Despite its simplicity, the algorithm is widely used in data classification and has been shown to produce results comparable to other classification methods in various domains.

7. Performance

The performance of Gradient Boosting and Naïve Bayes before and after feature selection will be evaluated and validated using the 10-fold cross-validation method and confusion matrix. Cross-validation is used to evaluate the performance of the prediction model by dividing the initial data into two parts, namely training and testing data, which were done 10 times. Although there is no fixed rule, the division of training and testing data is done with a ratio of 70:30 is often used in evaluating prediction models [20]. As for this study, eight performance metrics will be used, which include accuracy, balanced accuracy, sensitivity, specificity, precision, f-measure, G-means, and Area Under Curve (AUC).

Based on Table 1, some formulas can be calculated as follows:

(1) Accuracy

$$\text{Accuracy} = \frac{TP + TN}{TP + FP + FN + TN} \quad (4)$$

(2) Balanced accuracy

$$\text{Balanced accuracy} = \frac{(\text{Sensitivity} + \text{Specificity})}{2} \quad (5)$$

(3) Sensitivity

$$\text{Sensitivity} = \frac{TP}{TP + FN} \quad (6)$$

(4) Specivicity

$$\text{Specificity} = \frac{TN}{TN + FP} \quad (7)$$

(5) Precission

$$\text{Precision} = \frac{TP}{TP + FP} \quad (8)$$

(6) F – measure

$$F - \text{measure} = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (9)$$

(7) G – means

$$G - \text{means} = \sqrt{\text{Sensificity} \times \text{Specificity}} \quad (10)$$

(8) Area Under Curve (AUC)

$$AUC = \frac{TP + TN}{TP + FP + FN + TN} \quad (11)$$

A model with an AUC value close to 1 indicates that the model is very good, and if it is closer to 0, then it has the worst measurements [21]. To evaluate the performance of a binary classification model, it is essential to understand the relationship between actual and predicted class labels. One commonly used tool for this purpose is the confusion matrix, which summarizes the number of correct and incorrect predictions made by the classifier, broken down by each class. In the context of environmental data classification, for example, we might be interested in distinguishing between “Good” and “Polluted” categories. The confusion matrix not only provides insight into the overall accuracy of the model but also highlights the types of errors it makes—such as misclassifying polluted data as good, or vice versa. The structure of the confusion matrix for a binary classification problem is presented in Table 1.

Table 1. Confusion matrix for binary classification.

		Predicted	
Actual	Good	Good True Negative (TN)	Polluted False Positive (FP)
	Polluted	False Negative (FN)	True Positive (TP)

8. Methodology

8.1. Data Source

The data in this study are secondary data sourced from the Ministry of PUPR Directorate General of Water Resources of the Bengawan Solo River Basin Center, with a total of 720 data from January 2022 to December 2023.

8.2. Research Variables

The research variables used consist of 8 variables (parameters), namely pH (X_1), TDS (X_2), TSS (X_3), Temperature (X_4), BO (X_5), BOD (X_6), COD (X_7), NO₃ (X_8), and Water Quality Status (Y).

8.3. Research Procedure

The steps in this research are as follows:

- (1) Identify the problem.
- (2) Conduct a literature study.
- (3) Obtained secondary data from the Ministry of PUPR Directorate General of Water Resources Bengawan Solo River Basin Center.
- (4) Perform data description to describe and describe the data used.
- (5) At this stage, summary results are displayed to show the comparison of the number of polluted and unpolluted river classifications. In addition, data exploration analysis is carried out with simple descriptive statistics.
- (6) Perform pre-processing of data viz:
 - a. Cleaning data, making sure there are no *missing values*, and filling in the blanks with the *Expectation Maximization* algorithm.
 - b. Scaling the data, transforming the data using a standard scaler (Z Score).
- (7) Identify the imbalance ratio using a bar graph.
- (8) Data partitioning is performed by dividing the data into training and validation sets with a scheme of 70%

training data and 30% testing data.

- (9) Perform feature selection with the Wrapper method using the Boruta algorithm and Embedded method using random forest classifier.
- (10) Resampling data on training data using individual sampling, namely Rapidly Converging Gibbs Sampler (RACOG), and hybrid sampling, namely RACOG- RUS.
- (11) Perform classification using Gradient Boosting and Naïve Bayes algorithms.
- (12) Comparing classification performance with or without feature selection of RACOG and RACOG-RUS samples.
- (13) Draw conclusions.

Comparison and evaluation of model performance using confusion matrix obtained from the validation process using k-fold cross-validation with 10-fold cross-validation. The quality of the model can be seen based on the value of balanced accuracy, accuracy, specificity, sensitivity, precision, f-measure, recall, and Area Under Curve (AUC) for both training and validation sets.

Results and Discussion

9.1. Data Description

The description of river water quality parameter data can be seen in Table 2 below.

Table 2. Description of water quality parameters.

	Temperature	pH	Nitrate	DO	COD	KOB	TSS	TDS
Min	23.48	3.47	0.01	0.40	1.00	0.20	2.00	0.05
Max	36.26	9.24	18.90	51.00	495.80	59.50	4180.00	1223.00
Mean	27.54	6.61	0.97	5.33	30.58	4.41	91.22	869.44

Based on Table 2 shows descriptive statistics of water quality parameters. The lowest temperature parameter is 23.48, and the highest is 36.26, with an average river water temperature of 27.54. Furthermore, the lowest pH is 3.47, and the highest is 9.24, with an average pH in river water of 6.61. Furthermore, the lowest nitrate is 0.01, and the highest is 18.90, with an average nitrate in river water of 0.97. Furthermore, the lowest DO is 0.40, and the highest is 51.00, with an average DO in river water of 5.33. Furthermore, the lowest COD is 1.00, and the highest is 495.80, with an average COD in river water of 495.80. Furthermore, the lowest KOB is 0.20, and the highest is 59.50, with an average KOB in river water of 4.41. Furthermore, the lowest TSS is 2.00, and the highest is 4180.00, with an average TSS in river water of 91.22. Furthermore, the lowest TDS is 0.05, and the highest is 1223.00, with an average TDS in river water of 869.44.

9.2. Data Pre-Processing

Pre-processing of river water quality parameter data is shown in Table 3 below.

Table 3. Data pre-processing results.

N	Temperature	pH	Nitrate	DO	COD	KOB	TSS	TDS
1	-0.749	-0.319	-0.880	-1.060	-0.128	1.096	-0.272	-0.137
2	-0.543	-0.381	-0.890	-1.170	-0.420	1.490	-0.268	-0.137
3	-0.247	-1.708	-0.601	-0.866	-0.294	0.122	-0.247	-0.137
⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮
720	2.952	1.488	0.234	1.348	0.650	0.889	-0.197	6.095

Based on Table 3, pre-processing of river water quality has been carried out, and there are no variables or parameters that have a value of more than 2.5. So, it is concluded that there are no outliers in the research data.

9.3. Imbalanced Ratio

To find out how big the imbalance of a data set is, one must look at the imbalance ratio. The imbalance ratio is calculated by dividing the number of majority classes by the number of minority classes. When $IR = 1$, it can be said that the dataset is in a balanced position. In this study, it can be seen that the imbalance ratio is 71, so it can be said that the majority class is 71 times more than the minority class. The unbalanced classes between polluted and unpolluted classes are shown in Figure 1 below.

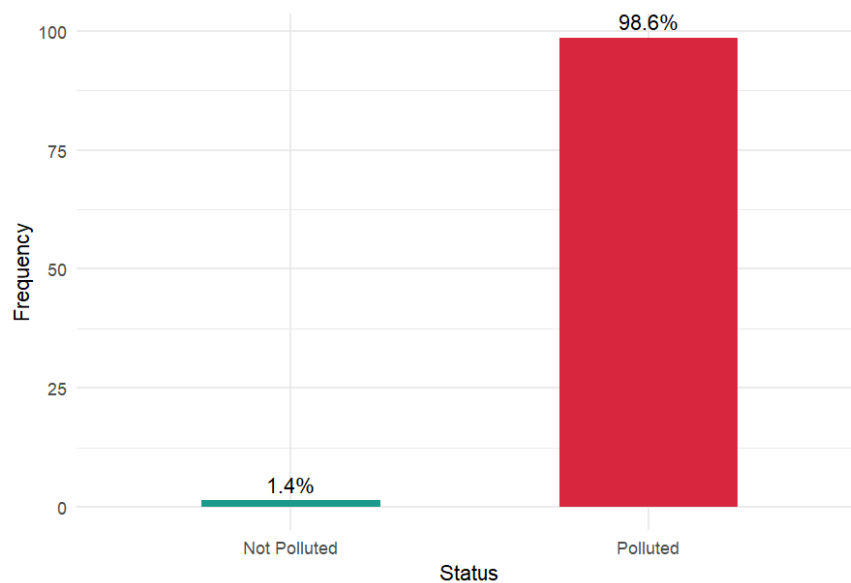


Figure 1. Dataset Imbalance Ratio.

9.4. Feature Selection

Feature selection serves to show the importance of features using the Boruta algorithm. Based on Figure 2, the boxplot for each feature shows that 5 influential and important attributes include Chemical Oxygen Demand (COD), Biochemical Oxygen Demand (BOD), Total Suspended Solids (TSS), potential of hydrogen, and Total Dissolved Solids (TDS). The embedded method is used to support the decision and to select features.

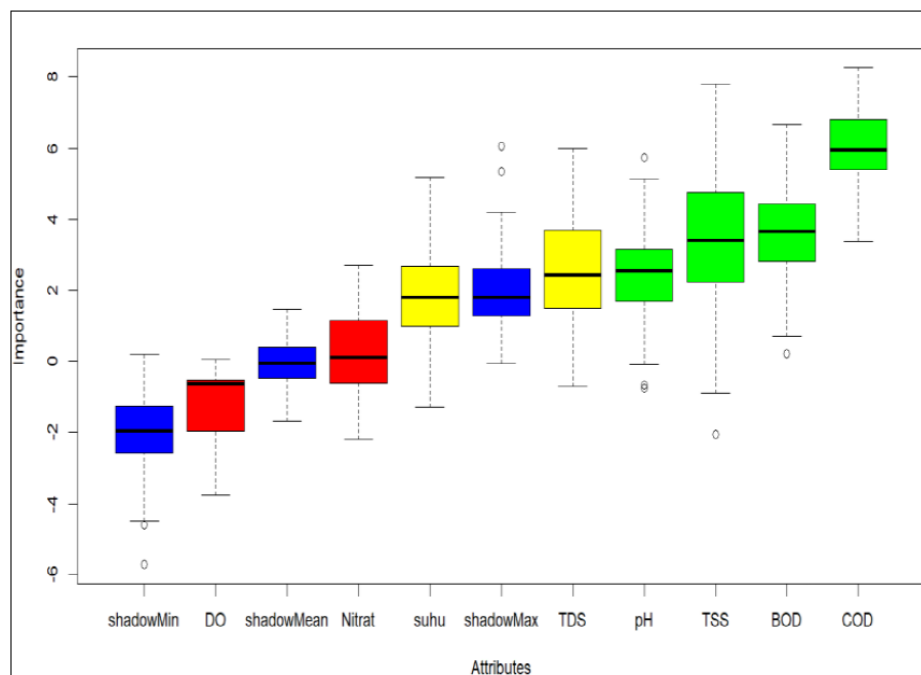


Figure 2. Feature Importance Using Boruta Algorithm.

This study uses Random Forest as a classification because it provides a simple method of feature selection using the mean decrease in accuracy and node cleanliness. When the average decrease in accuracy (gini index) is higher, the variable is considered more important in the model. The selected features are used at the top of the tree and contribute to the final decision to predict a larger proportion of the input samples. Furthermore, to find out the results of feature selection using the Variable Importance Plot for Random Forest from the Embedded algorithm is shown in Figure 3 below.

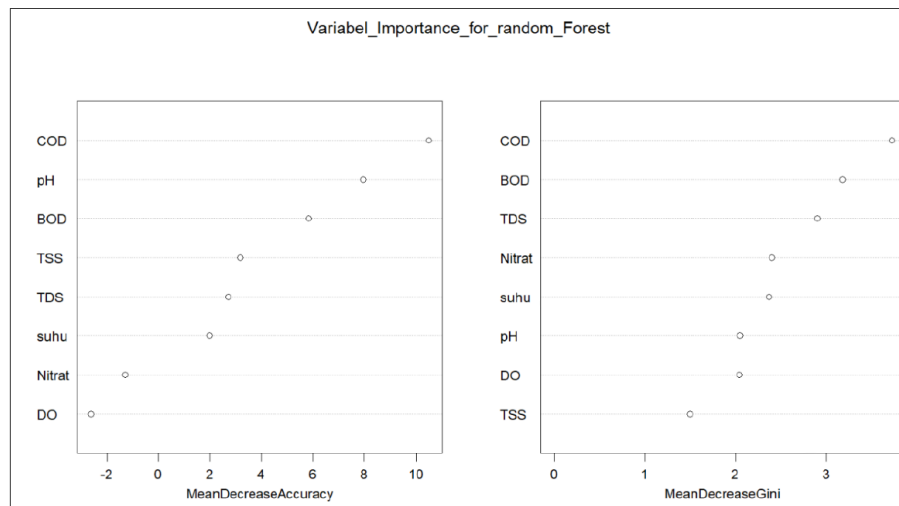


Figure 3. Feature Importance Using Embedded Algorithm.

Based on the output of Mean Decrease Accuracy, it is known that there are three features that are not important, including Dissolved Oxygen (DO), Nitrate (NO₃), and Temperature. This is in accordance with the wrapper method that has been used in previous research. Therefore, both Wrapper and Embedded methods show similar results.

9.5. Gradient Boosting

This section presents the performance of the Gradient Boosting ensemble model without and with feature selection using both the Boruta algorithm and the embedded method with a random forest classifier. In the performance comparison, the effect of individual sampling (RACOG) and hybrid sampling (RACOG-RUS) in classifying water quality is examined. Performance is evaluated using eight performance metrics: accuracy, balanced accuracy, sensitivity, specificity, precision, f-measure, and AUC. Furthermore, Table 4 presents the performance of Gradient Boosting with and without feature selection. In this table, the bolded values represent the highest performance scores in each scenario based on the number of instances and sampling method.

Table 4. Performance of Gradient Boosting with and without feature selection.

Algorithm	Feature Selection	Sampling	Number of Instances	Accuracy	Balanced Accuracy	Sensitivity	Spesificity	Precision	F-Measure	AUC
Gradient Boosting	With	RACOG	100	0.97	0.49	0.99	0.00	0.99	0.99	0.50
			200	0.98	0.50	0.99	0.00	0.99	0.99	0.67
			300	0.98	0.50	0.99	0.00	0.99	0.99	0.67
		RACOG-RUS	100	0.95	0.48	0.96	0.00	0.99	0.97	0.64
			200	0.97	0.49	0.98	0.00	0.99	0.98	0.66
			300	0.96	0.49	0.98	0.00	0.99	0.98	0.66
	Without	RACOG	100	0.98	0.50	0.99	0.00	0.99	0.99	0.50
			200	0.95	0.48	0.97	0.00	0.99	0.98	0.50
			300	0.96	0.49	0.98	0.00	0.99	0.98	0.50
		RACOG-RUS	100	0.90	0.46	0.92	0.00	0.98	0.95	0.65
			200	0.96	0.49	0.98	0.00	0.99	0.98	0.66
			300	0.93	0.47	0.94	0.00	0.99	0.96	0.50

Table 4 shows that RACOG outperforms RACOG-RUS from all performance metrics, except AUC, by using the number of numInstances of 100. Likewise, when feature selection is carried out, It can also be seen that RACOG outperforms RACOG-RUS across all performance metrics when the number of instances is 200.

9.6. Naïve Bayes

This section presents the performance of the Naïve Bayes model when not and with feature selection both with the Boruta algorithm and the embedded method. Furthermore, Table 5 presents the performance of Naïve Bayes with and without feature selection. In this table, the bolded values represent the highest performance scores in each scenario based on the number of instances and sampling method.

Table 5. Naïve Bayes performance with and without feature selection.

Algorithm	Feature Selection	Sampling	Number of Instances	Accuracy	Balanced Accuracy	Sensitivity	Spesificity	Precision	F-Measure	AUC
Naïve Bayes	With	RACOG	100	0.97	0.49	0.99	0.00	0.99	0.99	0.61
			200	0.98	0.50	0.99	0.00	0.99	0.99	0.60
			300	0.98	0.50	0.99	0.00	0.99	0.99	0.60
		RACOG-RUS	100	0.80	0.57	0.80	0.33	0.99	0.89	0.71
			200	0.84	0.58	0.84	0.33	0.99	0.91	0.74
			300	0.85	0.60	0.86	0.33	0.99	0.92	0.59
	Without	RACOG	100	0.87	0.61	0.89	0.33	0.99	0.94	0.49
			200	0.85	0.59	0.85	0.33	0.99	0.92	0.48
			300	0.84	0.59	0.84	0.33	0.99	0.91	0.48
		RACOG-RUS	100	0.73	0.70	0.73	0.67	0.99	0.84	0.47
			200	0.76	0.71	0.76	0.67	0.99	0.86	0.47
			300	0.82	0.58	0.83	0.33	0.99	0.90	0.48

Table 5 shows that RACOG outperforms RACOG-RUS in all performance metrics, except AUC, by using the number of numbers of 100. Likewise, when feature selection is carried out, it can also be seen that RACOG also outperforms RACOG-RUS from all performance metrics when the number of instances is 200.

9.7. Receiver Operating Characteristic (ROC) Plot of Gradient Boosting and Naïve Bayes Algorithms

Theoretically, if the classifier predicts the sample correctly, then the ROC curve will rise to the upper left of the graph. Furthermore, Figure 4 shows the ROC plot of each model from resampling RACOG and RACOG-RUS data on the Gradient Boosting classification algorithm with and without feature selection.

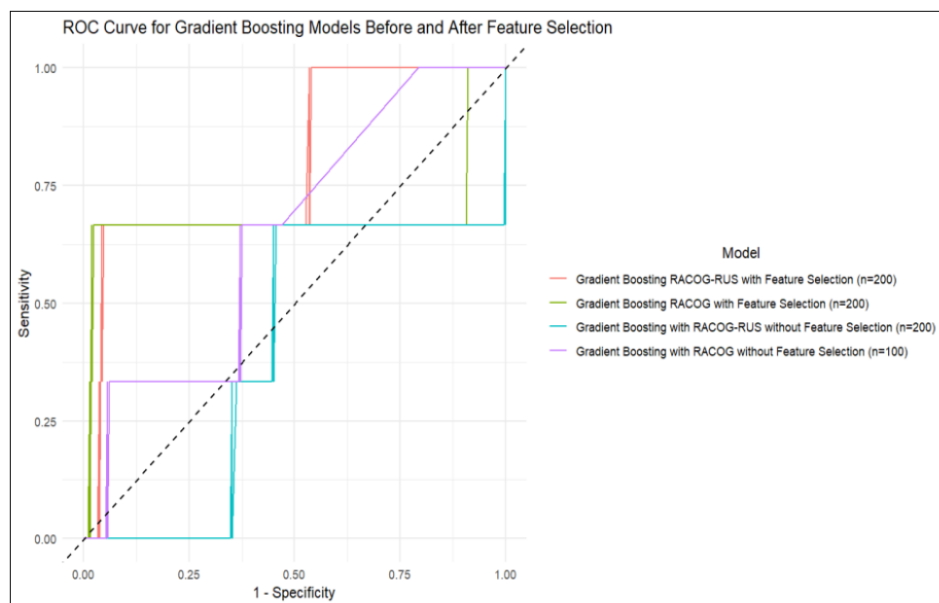
**Figure 4.** Comparison of RACOG and RACOG-RUS with and without Feature Selection on Gradient Boosting.

Figure 4 shows that Gradient Boosting correctly predicts samples with the highest performance when using RACOG-RUS resampling data with feature selection with a numIntances of 200. Furthermore, Figure 5 shows the ROC plot of each model from resampling RACOG and RACOG-RUS data in the Naïve Bayes classification algorithm with and without feature selection. Figure 5 shows that Naïve Bayes predicts samples correctly with the highest performance when using RACOG-RUS resampling data without feature selection with the number of numIntances 300.

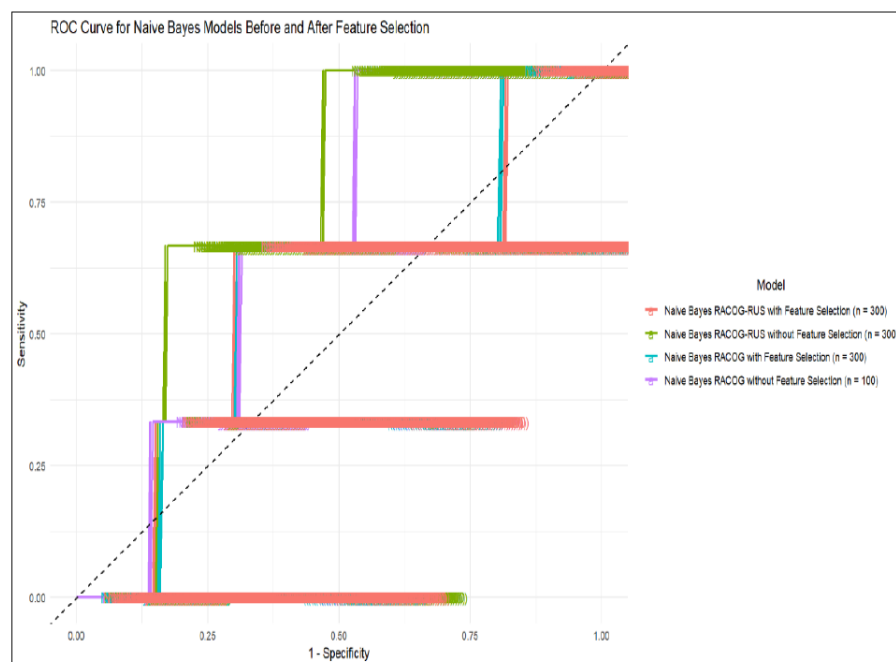


Figure 5. Comparison of RACOG and RACOG-RUS with and without Feature Selection on Naïve Bayes.

10. Conclusions

Based on the results and discussion, it is concluded that Gradient Boosting has the best performance when using RACOG-RUS resampling data with feature selection with a numInstances of 200. While Naïve Bayes has the best performance when using RACOG-RUS resampling data without feature selection with a numInstances of 300. It can be seen that resampling RACOG data does not outperform RACOG-RUS in both classification models because it is known that the data generated in RACOG does not make the dataset more balanced than RACOG-RUS, so it is necessary to do hybrid sampling if using RACOG samples as a training dataset.

Author Contributions

R.F.: conceptualization, methodology, software, writing—original draft preparation, visualization; H.K.: conceptualization, methodology, validation, supervision, writing—review and editing; D.D.P.: conceptualization, methodology, validation, supervision, writing—review and editing; D.A.S.: data curation, investigation, formal analysis; M.Y.M.: software support, validation, visualization. All authors have read and agreed to the published version of the manuscript.

Funding

This research was funded by the Indonesian Ministry of Education, Culture, Research, and Technology through the PMDSU Research Grant No. 038/E5/PG.02.00.PL/2024.

Institutional Review Board Statement

Not applicable.

Informed Consent Statement

Not applicable.

Data Availability Statement

The datasets generated and/or analyzed during the current study are available from the corresponding author on reasonable request. Due to privacy and ethical restrictions related to the data, public sharing is limited. However, the authors are committed to providing access to the data to qualified researchers upon reasonable request, ensuring compliance with confidentiality agreements and legal requirements.

Acknowledgments

The authors gratefully acknowledge the financial support from the Indonesian Ministry of Education, Culture, Research, and Technology through the PMDSU Research Grant No. 038/E5/PG.02.00.PL/2024. We also thank Institut Teknologi Sepuluh Nopember, Pattimura University, and the Ministry of Environment for their valuable support and collaboration throughout this research.

Conflicts of Interest

The authors declare no conflict of interest.

References

1. Khan, M.S.I.; Islam, N.; Uddin, J.; et al. Water quality prediction and classification based on principal component regression and gradient boosting classifier approach. *J. King Saud. Univ. Comput. Inf. Sci.* **2022**, *34*, 4773–4781. <https://doi.org/10.1016/j.jksuci.2021.06.003>.
2. Azhar, S.C.; Aris, A.Z.; Yusoff, M.K.; et al. Classification of River Water Quality Using Multivariate Analysis. *Procedia Environ. Sci.* **2015**, *30*, 79–84. <https://doi.org/10.1016/j.proenv.2015.10.014>.
3. Abo-Zahhad, M. Design of Smart Wearable System for Sleep Tracking Using SVM and Multi-Sensor Approach. *J. Eng. Sci.* **2023**, *51*, 1–15. <https://doi.org/10.21608/jesaun.2023.205964.1220>.
4. Hanisah, N.; Malek, A.; Fairos, W.; et al. Performance Evaluation of Classification Methods with Hybrid Sampling for Imbalanced Data: A Comparative Simulation Study. *SSRN* 2023. Available online: <https://ssrn.com/abstract=4519776> (accessed on 11 July 2023).
5. Ramyachitra, D.; Manikandan, P. Imbalanced Dataset Classification and Solutions: A Review. *Int. J. Comput. Bus. Res.* **2014**, *5*, 1–29.
6. Spelman, V.S.; Porkodi, R. A Review on Handling Imbalanced Data. In Proceedings of the International Conference on Current Trends towards Converging Technologies (ICCTCT), Coimbatore, India, 1–3 March 2018; pp. 1–11. <https://doi.org/10.1109/ICCTCT.2018.8551020>.
7. Das, B.; Krishnan, N.C.; Cook, D.J. RACOG and wRACOG: Two probabilistic oversampling techniques. *IEEE Trans. Knowl. Data Eng.* **2015**, *27*, 222–234. <https://doi.org/10.1109/TKDE.2014.2324567>.
8. Tyagi, S.; Mittal, S. Sampling approaches for imbalanced data classification problem in machine learning. In *Lecture Notes in Electrical Engineering*; Springer: Berlin/Heidelberg, Germany, 2020; pp. 209–221. https://doi.org/10.1007/978-3-030-29407-6_17.
9. Malek, N.H.A.; Yaacob, W.F.W.; Wah, Y.B.; et al. Comparison of ensemble hybrid sampling with bagging and boosting machine learning approach for imbalanced data. *Indones. J. Electr. Eng. Comput. Sci.* **2023**, *29*, 598–608. <https://doi.org/10.11591/ijeecs.v29.i1>.
10. Breiman, L. *Bagging Predictors*; Kluwer Academic Publishers: Dordrecht, The Netherlands, 1996.
11. Sahin, E.K. Assessing the predictive capability of ensemble tree methods for landslide susceptibility mapping using XGBoost, gradient boosting machine, and random forest. *SN Appl. Sci.* **2020**, *2*, 1308. <https://doi.org/10.1007/s42452-020-3060-1>.
12. Klug, M.; Barash, Y.; Bechler, S.; et al. A Gradient Boosting Machine Learning Model for Predicting Early Mortality in the Emergency Department Triage: Devising a Nine-Point Triage Score. *J. Gen. Intern. Med.* **2020**, *35*, 220–227. <https://doi.org/10.1007/s11606-019-05512-7>.
13. Mitchell, T.M. Does Machine Learning Really Work? 1997. Available online: www.kdnuggets.com/ (accessed on 11 July 2023).
14. Müller, A.C.; Guido, S. Introduction to Machine Learning with Python: A Guide for Data Scientists Introduction to Machine Learning with Python. O'Reilly Media, Inc.: Sebastopol, CA, USA.
15. Huan, Y.; Kong, Q.; Mou, H.; et al. Antimicrobial Peptides: Classification, Design, Application and Research Progress in Multiple Fields. *Front. Microbiol.* **2020**, *11*, 582779. <https://doi.org/10.3389/fmicb.2020.582779>.
16. van Ravenzwaaij, D.; Cassey, P.; Brown, S.D. A simple introduction to Markov Chain Monte–Carlo sampling. *Psychon. Bull. Rev.* **2018**, *25*, 143–154. <https://doi.org/10.3758/s13423-016-1015-8>.
17. Fadhilah, R.; Kuswanto, H.; Prastyo, D.D. Performance Analysis of Random Forest with Sampling for River Water Quality Classification. In Proceedings of the 2024 7th International Conference on Informatics and Computational Sciences (ICICoS), Semarang, Indonesia, 17–18 July 2024; pp. 456–461. <https://doi.org/10.1109/ICICoS62600.2024.10636858>.
18. Hanisah, N.; Malek, A.; Fairos, W.; et al. A New Hybrid Sampling for Classifying Imbalanced Data Based on Ensemble Decision Tree. 2023. Available online: <https://ssrn.com/abstract=4485808> (accessed on 11 July 2023).
19. Şahin, M. Impact of weather on COVID-19 pandemic in Turkey. *Sci. Total Environ.* **2020**, *728*, 138810. <https://doi.org/10.1016/j.scitotenv.2020.138810>.

20. Khosravi, Y.; Asilian-Mahabadi, H.; Hajizadeh, E.; et al. Factors influencing unsafe behaviors and accidents on construction sites: A review. *Int. J. Occup. Saf. Ergon.* **2014**, 20, 111–125. <https://doi.org/10.1080/10803548.2014.11077023>.
21. Liu, Y.; Li, Y.; Xie, D. Implications of imbalanced datasets for empirical ROC-AUC estimation in binary classification tasks. *J. Stat. Comput. Simul.* **2024**, 94, 183–203. <https://doi.org/10.1080/00949655.2023.2238235>.